

Nexus Concordat, Inc.

Founder: Marjorie Gayle McCubbins, BSc

caelsereith@aetherprotocols.com

EIN 39-2932385 SAM.gov UEI WJ3SQQ24JHR4

August 6, 2026

Dockets Management Staff (HFA-305)

Food and Drug Administration

5630 Fishers Lane, Rm. 1061

Rockville, Maryland 20852

Re: Public Comment on Docket No. FDA-2026-N-8563, Regulatory Science Innovations Catalyzing Medical Device Development; Co-sponsored Public Meeting between FDA CDRH and the Medical Device Innovation Consortium (MDIC), scheduled for September 25, 2026.

Dear Dockets Management Staff,

Nexus Concordat respectfully submits this public comment in response to the Agency's request for stakeholder perspectives on emerging device technologies, evolving clinical scenarios, and the critical regulatory science gaps that will shape medical device development over the next decade. We thank the Center for Devices and Radiological Health and the Medical Device Innovation Consortium for the opportunity to contribute to this long-term strategic planning effort.

Statement of interest.

Nexus Concordat is a Delaware C-Corporation (EIN 39-2932385, SAM.gov UEI WJ3SQQ24JHR4) developing a biology-native compiled artificial intelligence substrate for use in regulated decision environments, including Software as a Medical Device (SaMD) applications. The architecture is the subject of four pending United States provisional patent applications: USPTO 63/939,190 (Substrate-Independent Memory Weighting Architecture), USPTO 63/962,385 (Neurochemical Language Model), USPTO 63/988,485 (AETHER: Adaptive Endocrine Transformer Heads), and USPTO 64/034,536 (Scyla: Biologically-Native Programming Language). The operating subsidiary, NeXus Neuroinformatics, Inc., is preparing a Class II De Novo pathway submission for the Aegis clinical decision support

instrument.

Nexus Concordat has filed seven prior public comments in the Operation TrialBlazer policy family and adjacent FDA regulatory guidance dockets, including FDA-2026-N-4390 (AI-Enabled Optimization of Early-Phase Clinical Trials), FDA-2026-N-4699 (Expedited IND Pilot Program), FDA-2026-N-4492 (Drug Repurposing), FDA-2026-D-6539 (QSP-Based Dose Selection for MABEL), FDA-2025-D-6131 (New Approach Methodologies in Drug Development), and FDA-2024-D-4689 (Considerations for Biological Products, two comments). The perspectives shared in this comment build on the same architectural thesis presented in that prior body of work, applied here to the CDRH regulatory science planning horizon.

Executive summary.

We offer the Agency and MDIC three concrete recommendations for the next-decade regulatory science research portfolio, each developed in the sections that follow.

1. **Recognize Neurochemical Language Models (NLMs), and the architecturally-governed NLM subclass in particular, as an emerging device technology category distinct from statistically-trained AI/ML SaMD.** Current CDRH review pathways treat all AI/ML SaMD as members of a single technology category regulated primarily through post-hoc validation of outputs. Neurochemical Language Models are a distinct AI category (USPTO 63/962,385) that *mathematically* models human neurochemistry as a first-class computational element, with hormone-state vector dynamics, pharmacokinetic-analog decay equations, and trajectory-encoding operations implemented as formal mathematical structures rather than as biological metaphor or brain-inspired approximation. Within that category, NLMs whose governance mechanisms are compiler-enforced at the substrate level, producing inspectable reasoning artifacts and deterministic gate behavior by construction, represent a materially different regulatory subject and merit a distinct regulatory science framing.
2. **Prioritize research on clinician review of AI training corpus (inputs), not only AI outputs.** Every current regulatory framework in the field requires clinician oversight of AI outputs. Almost none require clinician oversight of AI inputs. This is not merely a structural gap in the current regulation floor; it is a structural failure. An AI system that has not been validated on what it was trained on has no authority to act in a regulated clinical context. Drugs derive their authority to act from validation of what is in the bottle. Clinical laboratory assays derive their authority to act from valida-

tion of the analyte, the method, and the reagents. AI systems in regulated clinical contexts should meet a comparable authority standard. Closing this failure requires research on training corpus manifests, per-row provenance, clinician-inspectable disclosure standards, and workflow integration for pre-deployment corpus review.

3. **Fund regulatory science research on the validation paradigm shift from output-validation to state-persistent reasoning pathway validation.** As AI systems capable of maintaining persistent internal state begin producing inspectable reasoning artifacts anchored in that state, the regulatory validation question shifts from “was the output correct” to “can the reasoning pathway be inspected, audited, and defended against the state history that produced it.” This is a materially different question from post-hoc output validation or from “explainable AI” techniques that generate justifications without persistent state. This shift will define the next decade of AI-enabled device regulation and requires dedicated research investment now.

1. Emerging device technology: architecturally-governed Neurochemical Language Models

The AI-enabled medical device landscape today is dominated by statistically-trained models, primarily standard large language models (LLMs) and specialized machine learning systems, whose outputs are validated post-hoc through comparison against reference standards. This is the technology profile current CDRH review pathways were designed to evaluate.

A distinct emerging technology category is now maturing: **Neurochemical Language Models (NLMs)**. Neurochemical Language Models are AI systems whose reasoning architecture *mathematically* models human neurochemistry as a first-class computational element. Hormone-state dynamics are implemented as vector mathematics with defined operations across ten neurochemical channels. Memory persistence follows pharmacokinetic-analog decay equations with substrate-independent decay constants derived from published biological half-lives. Trust tracking is implemented as multi-dimensional lattice operations with contextual weighting. These are formal mathematical structures, not biological metaphors, brain-inspired heuristics, or “neural-network-flavored” branding of statistically-trained models. As one concrete illustration: gated reasoning in NLMs can be implemented via softmax-weighted attention operating over Hodgkin-Huxley-style ion channel gating dynamics (Hodgkin and Huxley, 1952), a computational structure whose transitions are governed by defined differential equations grounded in classical neurophysiology rather than by learned probability weights alone. The NLM category is the subject of pending United States provisional patent USPTO 63/962,385 (Neurochemical Language Model), with the memory-weighting mathematics further specified in USPTO 63/939,190 (Substrate-Independent Memory Weighting Architecture). NLMs are architecturally different from standard LLMs; they are not fine-tuned LLMs, retrieval-augmented LLMs, or LLMs with an inference-time filter. They are a distinct category of AI system defined by their underlying mathematics.

Within the NLM category, **architecturally-governed NLMs** represent the specific subclass relevant to Software as a Medical Device deployment. In an architecturally-governed NLM, governance mechanisms (evidence gates, deterministic thresholds, inspectable reasoning state, cryptographically-signed provenance) are compiled into the substrate itself, at the compiler and runtime level, rather than layered on top of an otherwise-unconstrained model. In this subclass the AI system does not simply produce output; it earns the right to

produce output by traversing a defined sequence of gates, each of which either passes or fails against pre-committed evidence and integrity criteria. When a gate fails, the system does not guess forward; it refuses, requests additional evidence, or escalates to human judgment.

Architecturally-governed NLMs are not a variant of standard AI/ML SaMD. They represent a different regulatory subject because:

- **Reasoning is inspectable as an artifact.** The system exposes its own reasoning state, decision provenance, and evidence trail as a queryable object rather than as an internal computation that must be reverse-engineered from output behavior.
- **Governance is enforced by construction.** Compiler-enforced gate architecture means the system cannot produce output that bypasses the gate sequence. Governance is not a policy layered on top; governance is the substrate.
- **Behavior is deterministic under specified conditions.** Same inputs, same substrate version, same execution environment yield byte-identical output. This is a testable, auditable property that standard LLM outputs typically do not exhibit.
- **The underlying computational model is biology-anchored.** Hormone-state modulation, scar-memory persistence, and trust-lattice tracking are not metaphors; they are computational elements with defined mathematical operations and pharmacokinetic-analog decay constants (see USPTO 63/939,190, Substrate-Independent Memory Weighting Architecture).

The current CDRH review framework does not have a category for Neurochemical Language Models, either as a distinct AI technology class or as an architecturally-governed SaMD subclass. The evolving clinical scenarios in which architecturally-governed NLMs will be deployed (clinical decision support, individualized physiological state tracking in drug development, longitudinal measurement in regulated trials, patient-specific neurochemical state estimation) will not be well served if the framework treats them as ordinary AI/ML SaMD. The regulatory science research priority we recommend is the development of a distinct evaluation and qualification framework for the NLM category, and specifically for the architecturally-governed subclass proposed for SaMD use.

We note that the Agency already possesses a mature state-machine review methodology from prior work on deterministic embedded medical devices. The Pacemaker Formal Methods Challenge, initiated by the North American Software Certification Consortium

in 2007 in conjunction with release of the Pacemaker System Specification, and the Infusion Pump Software Safety Research program at FDA (FDA, 2024b), together with the substantial body of academic literature on formal verification of implantable and infusion medical device software (Jiang et al., 2012), collectively establish that state-machine formalisms are within the Agency’s regulatory-science methodological repertoire. Separately, the Agency’s AI/ML SaMD framework treats machine learning models as either “locked” at deployment or “adaptive” under a Predetermined Change Control Plan (FDA, 2024a; FDA, 2021b; FDA, 2025). Neither of those AI frames captures a state-persistent AI architecture, in which the model itself remains fixed while a defined mathematical internal state evolves per query and per session. The bridge between the two, extending state-machine formal review methodology to state-persistent AI architectures, is a specific regulatory science gap we recommend the Agency include in the next-decade research portfolio.

2. Critical regulatory science failure: clinician review of AI inputs, not only outputs

Every regulatory framework currently anchoring AI in healthcare requires clinician oversight of AI outputs. NIST AI Risk Management Framework 1.0 (NIST, 2023), ISO/IEC 42001 (ISO/IEC, 2023), FDA Good Machine Learning Practice (FDA, Health Canada, and MHRA, 2021; IMDRF, 2025), FDA Predetermined Change Control Plan (FDA, 2024a), and the World Health Organization principles on AI for health (WHO, 2021; WHO, 2023) all frame the human-in-the-loop question as “how does the clinician evaluate what the AI produces.” Almost none require the clinician to evaluate what the AI was trained on.

This is not merely a structural blind spot. It is a structural failure of the current regulatory approach to derive legitimate authority for an AI system to act in a regulated clinical context. Every other category of regulated instrument that acts on patient care earns its authority to act through pre-market validation of its inputs, not only its outputs. A drug is approved because the Agency validates what is in the bottle. A clinical laboratory assay is approved because the Agency validates the analyte, the method, and the reagents. A medical device is approved because the Agency validates its design inputs, its manufacturing process, and its intended-use specification. In each case authority to act on patient care is derived from pre-market validation of what the instrument *is* and *contains*, not only what it *produces*.

An AI system deployed in a regulated clinical context should meet a comparable authority standard. An AI system that has not been validated on what it was trained on has not established legitimate authority to act. Post-hoc output validation is a necessary condition; it is not a sufficient one. A clinician who reviews only the output of a clinical AI can validate whether a single recommendation was correct once. A clinician who reviews the training corpus can validate whether the model was ever equipped to answer the question at all. These are not the same review, and they do not substitute for each other. A clinician cannot answer whether an AI system is safe for the patient population in front of her without seeing what population the AI was trained on.

Closing this gap requires regulatory science research on:

- Training corpus manifest standards: what disclosures are required (row counts per source, license per source, capture date range, demographic representation summary) for the manifest to be meaningfully reviewable by a clinical audience.
- Per-row provenance requirements: standards for training data provenance sufficient to permit clinician inspection at the row level, not merely at the source-family level.
- Clinician review workflows: what pre-deployment and post-retraining review protocols are practical for the treating clinician (as opposed to the informatics or QA function).
- Vendor disclosure obligations: whether “vendor trade secret” can be recognized as a valid ground for withholding training corpus content from clinical review in a regulated deployment.
- Population representation review: what standard is required for comparing the deployed patient population to the training distribution, and what constitutes material divergence.

We recommend that CDRH include clinician review of AI training inputs as an explicit topic in the next decade regulatory science research portfolio.

3. Critical regulatory science gap: from output validation to state-persistent reasoning pathway validation

The dominant validation paradigm for AI-enabled devices today asks: does the output match a reference standard? This paradigm is well developed for narrow AI systems that produce discrete, comparable outputs (a classification, a segmentation, a numeric prediction). It becomes increasingly difficult to apply as AI systems produce composite outputs, generative content, or recommendations whose “correctness” is not defined against a single reference.

As architecturally-governed NLMs and adjacent state-persistent architectures mature, an alternative validation paradigm becomes available: validate the state-persistent reasoning pathway itself. Rather than asking whether every possible output is correct, evaluate whether the system’s decision-making architecture is inspectable, deterministic, and defensible against the internal state history that produced the decision. Ask whether every recommendation carries an evidence trail. Ask whether every gate the system traverses has a defined threshold and a documented pass/fail criterion. Ask whether the system can, on demand, produce a reproducible reasoning record, anchored in the state that existed at inference time, that a regulator or clinician can audit.

Persistent internal state is what distinguishes this validation paradigm from prior work on “explainable AI,” “AI transparency,” or chain-of-thought techniques. Explainability tools generate a post-hoc justification for a specific output; they do not require the AI system to maintain internal state that persists across queries or across time. Retrieval-augmented generation retrieves context freshly on each call; the retrieval trace is not persistent state. Chain-of-thought exposes intermediate reasoning tokens but does not persist reasoning across sessions. State-persistent architectures, Neurochemical Language Models being one example, maintain hormone-state, memory, and trust-context as first-class computational elements that carry forward between decisions.

This is not a semantic distinction and it is not a labeling choice. It is a mathematical one. In state-persistent architectures such as NLMs, the internal state is a mathematical object: a hormone-state vector across defined neurochemical channels, a scar-weighted memory structure with pharmacokinetic-analog decay operations, a multi-dimensional trust lattice with contextual weighting. State persistence means these mathematical objects carry forward through defined mathematical operations rather than being reconstructed or re-inferred from prompt context on each call. Validating a state-persistent reasoning

pathway therefore means validating a sequence of mathematical operations on mathematical objects. It is measurable, testable, and reproducible in the same sense that any other mathematical model in regulated science is measurable, testable, and reproducible. Reasoning pathway validation for state-persistent architectures evaluates not just “was this decision defensible in isolation” but “is the trajectory of decisions defensible against the state trajectory that produced them, and does that state trajectory itself conform to the mathematical model the system is built on.”

We recommend that CDRH include the following as regulatory science research priorities:

- Evaluation frameworks for state-persistent reasoning pathway artifacts: what constitutes an acceptable inspectable reasoning artifact, what state-history disclosures it must carry, how state anchoring is verified.
- Reproducibility standards for state-anchored AI recommendations: cryptographic signing, deterministic execution requirements, signed inference manifest emission that includes both the recommendation and the internal-state snapshot that produced it.
- Audit trail specifications: model version, input feature snapshot at inference time, internal state at inference time, output, confidence score, human reviewer identity, override or accept status, timestamp. Retention standards. Access controls at the tier appropriate for protected health information.
- Pre-committed threshold discipline including state-drift monitoring: standards for what performance thresholds (AUC floors, sensitivity floors for false-negative-critical use cases, calibration thresholds, state-drift monitoring cadence, retirement criteria) must be committed in writing before deployment.

A concrete example illustrates the paradigm. In an architecturally-governed NLM, the internal cortisol channel spikes mathematically in response to ambiguous evidence and elevated-stakes clinical context. That cortisol spike propagates through the substrate as a modulator of gate behavior, biasing the system toward escalation, refusal, or additional-evidence requests rather than confident output. When the system produces a recommendation, or when it refuses to, that decision is traceable back to a specific hormone-state trajectory. The mathematical justification is inspectable at every step: input evidence combined with stakes assessment produced this cortisol trajectory, which drove this gate posture, which produced this output. The output is a reasoned output whose reasoning is anchored in mathematical state, not in a post-hoc narrative reconstruction of what the

system “meant” to do. The higher the potential consequence of the decision, the more conservatively the substrate is mathematically calibrated to route it. This is exactly the property regulated-decision environments require of any instrument that participates in patient care.

A model that cannot be retired is not governed. A recommendation that cannot be inspected is not auditable. A system that cannot produce its reasoning artifact and its state snapshot on demand cannot be defended in a regulated environment. These are not aspirational goals. They are the mechanical requirements of an emerging validation paradigm the Agency’s next-decade research portfolio should anticipate.

4. Alignment with adjacent regulatory science instruments

The recommendations above are consistent with, and would strengthen, several instruments the Agency and adjacent bodies have already established:

- NIST AI Risk Management Framework 1.0 (NIST, 2023), particularly the Map and Measure functions applied at the substrate rather than post-hoc.
- ISO/IEC 42001:2023 organization-wide AI management system requirements (ISO/IEC, 2023), applied to the substrate design layer.
- FDA Good Machine Learning Practice Guiding Principles (FDA, Health Canada, and MHRA, 2021; adopted in final form by IMDRF, 2025), particularly the principles addressing data quality, representative training data, and model transparency, extended to require clinician-inspectable disclosure.
- FDA Predetermined Change Control Plan Final Guidance (FDA, 2024a), together with the FDA Draft Guidance on Artificial Intelligence-Enabled Device Software Functions Lifecycle Management (FDA, 2025), extended to recognize architecturally-enforced change control as a distinct category from procedurally-committed change control.
- World Health Organization Ethics and Governance of Artificial Intelligence for Health (WHO, 2021) and Regulatory Considerations on Artificial Intelligence for Health (WHO, 2023), particularly the requirement for patient voice in governance and clinician oversight of AI decisions.

- FDA Clinical Decision Support Software Final Guidance (FDA, 2022), extended to distinguish output-facing decision support from state-anchored reasoning support.

We further note that the framework this comment describes is presented in fuller form at nexusneurodata.com/regulation, a publicly available document with complete citation to the current regulation floor and adjacent instruments.

5. Closing framing question

We offer one framing question the Agency's regulatory science planning effort should hold as a load-bearing consideration through the next decade:

What kind of validation should AI require before we trust it with regulated decisions? Do we validate outputs? Do we validate performance metrics? Do we validate reasoning pathways? Or do we need a combination of all three?

We do not believe the field has answered this question yet. The Sept 25 public meeting, and the research priorities that emerge from it, present the opportunity to begin answering it deliberately rather than by accretion.

Closing

Nexus Concordat and its operating subsidiary NeXus Neuroinformatics support the Agency and MDIC in the long-term regulatory science planning effort this docket initiates. We would be pleased to elaborate on any of the recommendations above, to describe the operational mechanism of the architecturally-governed substrate in more technical detail, or to participate in follow-up dialogue that supports the Sept 25 public meeting and the next-decade research portfolio.

Respectfully submitted,

Marjorie Gayle McCubbins, BSc
Founder and Chief Executive Officer

Nexus Concordat, Inc.

caelsereith@aetherprotocols.com

References

All references below are hyperlinked to their primary source (agency site, DOI, or PubMed).

1. National Institute of Standards and Technology (2023). *Artificial Intelligence Risk Management Framework (AI RMF 1.0)*. NIST AI 100-1. doi.org/10.6028/NIST.AI.100-1
2. ISO/IEC (2023). *ISO/IEC 42001:2023. Information technology, Artificial intelligence, Management system*. iso.org/standard/81230.html
3. U.S. Food and Drug Administration, Health Canada, and Medicines and Healthcare products Regulatory Agency (2021). *Good Machine Learning Practice for Medical Device Development: Guiding Principles*. fda.gov
4. International Medical Device Regulators Forum (2025). *Good Machine Learning Practice for Medical Device Development: Guiding Principles*. Final adoption, January 2025. imdrf.org
5. U.S. Food and Drug Administration (2024a). *Marketing Submission Recommendations for a Predetermined Change Control Plan for Artificial Intelligence-Enabled Device Software Functions*. Final guidance. fda.gov
6. U.S. Food and Drug Administration (2024b). *Infusion Pump Software Safety Research at FDA*. fda.gov/medical-devices/infusion-pumps
7. U.S. Food and Drug Administration (2025). *Draft Guidance: Artificial Intelligence-Enabled Device Software Functions: Lifecycle Management and Marketing Submission Recommendations*. January 6, 2025. fda.gov
8. U.S. Food and Drug Administration (2022). *Clinical Decision Support Software: Final Guidance for Industry and FDA Staff*. fda.gov
9. U.S. Food and Drug Administration (2021b). *Artificial Intelligence and Machine Learning (AI/ML) Software as a Medical Device Action Plan*. fda.gov

10. World Health Organization (2021). *Ethics and Governance of Artificial Intelligence for Health: WHO guidance*. [who.int](https://www.who.int)
11. World Health Organization (2023). *Regulatory Considerations on Artificial Intelligence for Health*. [who.int](https://www.who.int)
12. Hodgkin, A. L., & Huxley, A. F. (1952). A quantitative description of membrane current and its application to conduction and excitation in nerve. *Journal of Physiology*, 117(4), 500–544. [PubMed 12991237](https://pubmed.ncbi.nlm.nih.gov/12991237/)
13. Jiang, Z., Pajic, M., Moarref, S., Alur, R., & Mangharam, R. (2012). Modeling and verification of a dual chamber implantable pacemaker. In *Tools and Algorithms for the Construction and Analysis of Systems (TACAS 2012)*, Lecture Notes in Computer Science 7214, Springer. doi.org/10.1007/978-3-642-28756-5_14
14. United States Patent and Trademark Office. *Provisional Patent Application 63/939,190: Substrate-Independent Memory Weighting Architecture*. Filed December 12, 2025 by Nexus Concordat, Inc.
15. United States Patent and Trademark Office. *Provisional Patent Application 63/962,385: Neurochemical Language Model*. Filed January 17, 2026 by Nexus Concordat, Inc.
16. United States Patent and Trademark Office. *Provisional Patent Application 63/988,485: AETHER: Adaptive Endocrine Transformer Heads*. Filed February 23, 2026 by Nexus Concordat, Inc.
17. United States Patent and Trademark Office. *Provisional Patent Application 64/034,536: Scyla: Biologically-Native Programming Language*. Filed April 9, 2026 by Nexus Concordat, Inc.

Prior Nexus Concordat public comments referenced in this submission are on the record at [regulations.gov](https://www.regulations.gov) under docket numbers [FDA-2026-N-4390](#), [FDA-2026-N-4699](#), [FDA-2026-N-4492](#), [FDA-2026-D-6539](#), [FDA-2025-D-6131](#), and [FDA-2024-D-4689](#).